

The Ethical and Legal Implications of Shadow AI in Sensitive Industries: A Focus on Healthcare, Finance, and Education

Abstract

This study examines the ethical and legal implications of Shadow AI in healthcare, finance, and education by analyzing unauthorized AI deployments and their impact on data privacy, cybersecurity, and regulatory compliance. Using a quantitative research approach, descriptive statistics, ordinal regression modeling, and network analysis were employed to assess AI violations using the MITRE ATLAS AI Incident Database, EU AI Act Public Database, and IBM X-Force Threat Intelligence Report. Findings reveal that privacy breaches are most prevalent in education (22 cases), bias-related issues dominate finance (20 cases), and cybersecurity risks are highest in healthcare (19 cases). Legal risk analysis shows a 20% probability of regulatory intervention, with breach type as the strongest determinant. Anomaly detection identified healthcare as the most vulnerable to AI-driven cyber threats (8 anomalies). Recommendations include sector-specific AI compliance frameworks, AI ethics committees, cybersecurity policies, and stricter regulatory enforcement.

Keywords: Shadow AI, AI governance, cybersecurity risks, regulatory compliance, algorithmic bias.

1. Introduction

The rapid integration of artificial intelligence (AI) across critical sectors such as healthcare, finance, and education has significantly enhanced operational efficiency, decision-making, and automation. However, alongside these sanctioned AI applications, an emerging trend known as Shadow AI has introduced considerable ethical and legal challenges (Sabin, 2025). Shadow AI refers to AI-driven tools and models implemented without regulatory oversight, often adopted by employees, third-party vendors, or decentralized teams without adhering to security protocols and compliance frameworks. Unlike AI systems governed by corporate IT policies, Shadow AI operates autonomously, bypassing essential scrutiny and exposing organizations to security vulnerabilities and ethical breaches (Editorial Team, 2025).

Unauthorized AI implementation in sensitive industries presents substantial risks. In healthcare, unregulated AI diagnostic tools may lead to patient confidentiality breaches,

misdiagnoses, and malpractice, exposing institutions to legal liability (Isibor, 2024). In finance, unchecked AI-driven models amplify the risk of algorithmic biases, unauthorized trading, and cyber threats, destabilizing financial markets (Zekos, 2021). Within education, unauthorized AI applications contribute to academic dishonesty, flawed grading mechanisms, and violations of student privacy, challenging institutional integrity. As AI adoption expands, the proliferation of Shadow AI demands urgent intervention from regulators, policymakers, and industry leaders to prevent widespread ethical and legal infractions.

Shadow AI has become increasingly prevalent due to the absence of formal AI governance strategies, prompting employees to seek third-party AI tools to enhance productivity. While these tools may streamline workflows, their unauthorized use introduces compliance violations, security vulnerabilities, and biased decision-making. The World Economic Forum's 2025 report on AI and cybersecurity highlights the risks of unregulated AI applications, emphasizing their role in cybersecurity threats and ethical dilemmas (World Economic Forum, 2025). Similarly, Wheeler (2025) identifies unauthorized AI applications in government agencies as a significant security risk, exposing regulatory gaps. These findings underscore the need for stringent AI governance measures in industries where ethical responsibility and legal compliance are critical.

Among the most pressing concerns surrounding Shadow AI is its impact on data privacy and security. In industries managing sensitive information—such as patient health records, financial transactions, and student academic data—unauthorized AI applications significantly increase confidentiality breach risks (Kessem, 2024; Schueler, 2024). IBM's 2024 Cost of a Data Breach Report indicates that unmonitored data sources linked to Shadow AI delay breach detection times and elevate security risks (IBM, 2024). The healthcare sector, in particular, faces rising costs, with the average breach reaching \$6.45 million (Seh et al., 2020). Between 2015 and 2019, hacking and IT-related breaches accounted for over 90% of exposed healthcare records, underscoring the security vulnerabilities posed by unauthorized AI (Davis, 2022).

Beyond data security, Shadow AI exacerbates algorithmic bias and discrimination. AI models deployed without proper oversight often lack fairness assessments, leading to biased decision-making that disproportionately affects marginalized groups. In finance, unregulated AI-driven credit scoring systems have resulted in discriminatory lending practices, prompting regulatory scrutiny (Zekos, 2021). In education, unauthorized AI grading algorithms have unfairly penalized non-native English speakers, leading to lawsuits and institutional policy revisions (Silor, 2024). The unchecked expansion of Shadow AI risks entrenching systemic biases, further marginalizing vulnerable populations and eroding public trust in AI-driven decision-making (Schueler, 2024).

The lack of accountability and transparency in Shadow AI deployment presents additional legal and ethical dilemmas. When AI-generated errors occur, determining liability becomes increasingly complex. In healthcare, for instance, an unauthorized AI-powered diagnostic tool at a U.S. hospital resulted in a misdiagnosis, prompting legal action and regulatory scrutiny (Abbas, 2024). Without clear governance frameworks, it remains unclear whether responsibility lies with the AI model, the employee who deployed it, or the institution that permitted its use. A similar case in finance emerged in 2025, when the U.S. Securities and Exchange Commission (SEC) fined a major bank for using an unauthorized AI-powered credit assessment tool that led to biased lending decisions and regulatory violations (Reuters, 2025).

Another area of concern is Shadow AI's role in AI-generated misinformation and deepfake threats. The ability of unregulated AI to produce synthetic content presents challenges across education, finance, and healthcare. In academia, AI-generated essays and deepfake student submissions have contributed to rising academic fraud. According to Freeman (2024) over 50% of students use AI for assignments, with 5% admitting to using AI-generated content for academic dishonesty. Similarly, in finance, unauthorized AI trading models have manipulated stock predictions and engaged in fraudulent insider trading, raising concerns about AI-powered economic manipulation (Zekos, 2021). In healthcare, AI-generated fake medical records have introduced risks related to fraudulent clinical trial data, potentially distorting public health policies and medical research (Isibor, 2024), whereas it is perceived that AI-generated misinformation in financial markets could lead to unprecedented economic manipulation, further emphasizing the risks associated with Shadow AI.

Real-world case studies illustrate the tangible consequences of Shadow AI (Abbas, 2024; Reuters, 2025; Walsh, 2020). In healthcare, the unauthorized use of an AI diagnostic tool at a U.S. hospital led to a misdiagnosis, prompting legal action and regulatory investigations. In finance, a 2025 SEC investigation found that a bank had implemented an AI-driven credit evaluation system without regulatory approval, leading to discriminatory lending practices and financial penalties. In education, a major UK university faced public backlash after its AI grading system disproportionately downgraded international students, resulting in legal challenges and policy reforms.

To address these risks, organizations must establish comprehensive AI governance frameworks incorporating compliance audits, AI ethics boards, and regulatory enforcement mechanisms. Governance policies should mandate regular audits of AI systems to ensure adherence to ethical and legal standards. Additionally, the implementation of AI ethics committees within organizations can provide oversight for AI deployment, ensuring that AI-driven decision-making aligns with principles of fairness, transparency, and accountability. Regulatory agencies, such as the SEC, the FDA, and

data protection authorities overseeing the General Data Protection Regulation, must strengthen their oversight of AI applications by imposing stricter penalties for unauthorized AI usage (Shandilya et al., 2024). This research aims to critically examine the ethical and legal implications of Shadow AI in sensitive industries—healthcare, finance, and education—by analyzing its impact on data privacy, decision-making integrity, security vulnerabilities, and regulatory compliance while proposing governance frameworks to mitigate associated risks, by achieving the following objectives:

1. Investigates the prevalence and impact of Shadow AI in healthcare, finance, and education, focusing on ethical dilemmas including data privacy violations, bias in decision-making, and accountability concerns.
2. Analyzes the legal and regulatory challenges posed by Shadow AI, assessing existing frameworks (GDPR, HIPAA, SEC, FERPA) and identifying gaps in AI governance and compliance.
3. Evaluates the role of Shadow AI in cybersecurity risks, AI-driven misinformation, and adversarial threats, particularly in deepfake content, synthetic financial fraud, and AI-enabled academic dishonesty.
4. Recommends effective strategies and governance models for mitigating Shadow AI risks, including AI ethics boards, regulatory enforcement mechanisms, and AI compliance auditing in organizations.

2. Literature Review

The rise of Shadow AI, or unauthorized AI applications, presents serious ethical challenges across critical sectors such as healthcare, finance, and education (Sabin, 2025). These systems operate outside regulatory oversight, raising concerns about data privacy, algorithmic bias, transparency, and misinformation (Sahota, 2024; Balogun et al., 2025). Without proper governance, Shadow AI exacerbates ethical risks, making intervention necessary to prevent harmful consequences (Sabin, 2025; Fabuyi et al., 2024).

A primary ethical concern is unauthorized access to sensitive data. AI models require extensive data to function effectively, yet when deployed without formal oversight, they pose security risks (Dhirani et al., 2023; Obioha-Val et al., 2025). In healthcare, unregulated AI systems compromise patient confidentiality, as legal protections often fail to safeguard clinical data against cyber threats (Isibor, 2024; Obioha-Val et al., 2025). Similarly, in finance, AI-driven tools processing transactions without authorization create opportunities for fraud and identity theft (Zekos, 2021; Kolade et al., 2025). In education, unauthorized AI applications can violate student privacy, increasing the risk of data breaches (Yang & Beil, 2024; Val et al., 2024). The absence of security protocols

in Shadow AI deployments heightens these vulnerabilities, underscoring the ethical responsibility of organizations to implement robust AI governance frameworks (Kurian, 2025; Obioha-Val et al., 2025).

Beyond privacy concerns, algorithmic bias remains a significant issue, particularly when AI models reinforce systemic disparities. In healthcare, AI diagnostic tools trained on biased datasets may misinterpret symptoms based on race or gender, exacerbating inequalities in medical treatment (Norori et al., 2021; Alao et al., 2024). Similarly, in finance, AI-driven credit scoring systems have demonstrated discriminatory lending practices, disproportionately affecting marginalized communities (Nuka & Osedahunsi, 2024; Joeaneke et al., 2024). The unregulated nature of Shadow AI heightens the risk of biased decision-making, as these systems bypass fairness assessments and regulatory scrutiny (Kurian, 2025; Arigbabu et al., 2024). Addressing bias requires diverse and representative training data, rigorous validation processes, and continuous monitoring to ensure equitable outcomes.

The lack of transparency in AI decision-making compounds these ethical concerns. Many AI systems, particularly those operating without oversight, function as "black boxes," making it difficult to determine how decisions are made (Hassija et al., 2023; Samuel-Okon et al., 2024). This opacity raises accountability questions, particularly when AI-driven decisions result in harm. Establishing responsibility—whether with developers, unauthorized users, or organizations failing to prevent unregulated AI deployment—remains a challenge (Habbal et al., 2024; Gbadebo et al., 2024). Ethical AI governance must prioritize clear documentation of AI decision-making and establish mechanisms for addressing unintended harms.

The spread of misinformation is another pressing issue linked to Shadow AI. AI-generated content often appears credible yet may be factually incorrect, posing risks in healthcare, finance, and education (Kurian, 2025). Inaccurate medical information endangers patient safety, misleading financial data disrupts markets, and AI-generated academic content threatens institutional integrity (Bala et al., 2024; Olabanji et al., 2024). Addressing these risks requires the implementation of detection technologies, ethical AI guidelines, and stronger regulatory enforcement.

Legal and Regulatory Challenges in Addressing Shadow AI

The rapid advancement of artificial intelligence (AI) has outpaced the development of comprehensive legal and regulatory frameworks, creating significant challenges in managing unauthorized AI applications, commonly referred to as Shadow AI. While regulations such as the European Union's General Data Protection Regulation (GDPR), the Health Insurance Portability and Accountability Act (HIPAA) in the United States, and the Family Educational Rights and Privacy Act (FERPA) provide foundational AI

governance structures, they fail to address the complexities associated with Shadow AI (European Commission, 2023). These regulations primarily focus on data protection, algorithmic fairness, and transparency but lack provisions for AI systems operating outside formal oversight, leaving a regulatory gap that enables Shadow AI to proliferate.

One of the most pressing legal challenges is the detection and prevention of unauthorized AI deployments. Existing regulatory frameworks primarily target sanctioned AI applications, leaving a blind spot for systems implemented without approval (Walter, 2024; John-Otumu et al., 2024). This challenge is particularly concerning in healthcare, finance, and education, where Shadow AI can lead to data breaches, algorithmic biases, and unethical decision-making. The ambiguity surrounding AI compliance further complicates enforcement efforts. While the GDPR enforces strict data protection rules, it does not explicitly regulate AI-specific concerns such as algorithmic transparency and accountability, allowing Shadow AI to operate largely unchecked (Huang et al., 2024; Salako et al., 2024). A recent case involving the Chinese AI startup DeepSeek highlights this regulatory gap, as Italy's data protection authority, Garante, imposed restrictions on the company after it failed to comply with privacy requirements (Reuters, 2025). This incident explains the difficulties regulators face in enforcing existing laws against unauthorized AI applications.

Legal liability for AI-driven failures presents another regulatory challenge. Traditional legal frameworks struggle to address the autonomous nature of AI decision-making, making it difficult to determine accountability when AI-generated errors cause harm (Akpuokwe et al., 2024; Joseph, 2024). In healthcare, misdiagnoses from unauthorized AI diagnostic tools raise questions about liability, while in finance, AI-driven discriminatory lending practices create legal and ethical concerns (Reuters, 2025). The decentralized nature of Shadow AI complicates accountability, as responsibility may fall on AI developers, unauthorized users, or organizations failing to enforce governance measures (Chin et al., 2024; Kolade et al., 2024). Courts and policymakers have begun addressing these concerns, as demonstrated by legal disputes such as Getty Images' lawsuit against Stability AI, which is accused of using copyrighted images to train its AI model without authorization (Brittain, 2023; Olateju et al., 2024). The outcome of this case, expected in 2025, is anticipated to shape AI governance and copyright law.

As AI-driven automation expands, regulators struggle to develop effective oversight mechanisms, particularly for decentralized AI deployments (Balakrishnan, 2024; Adigwe et al., 2024). AI models trained and deployed outside centralized control present significant monitoring challenges, further exposing organizations to legal risks (Habbal et al., 2024; Olabanji et al., 2024). Companies that fail to implement governance frameworks and actively monitor Shadow AI usage face substantial compliance violations and reputational damage.

Shadow AI and Cybersecurity Risks

The proliferation of unauthorized artificial intelligence applications, commonly referred to as Shadow AI, presents significant cybersecurity risks across multiple sectors. These unregulated AI tools operate outside established IT governance frameworks, increasing organizational vulnerabilities and exposing systems to malicious exploitation (Chernousov, 2024; Olabanji et al., 2024). Without the strict security protocols applied to officially sanctioned AI, Shadow AI expands the attack surface for cybercriminals, making it a critical concern for cybersecurity professionals (Editorial Team, 2025). Malik et al. (2024) underscores the susceptibility of such systems to adversarial attacks, wherein hackers manipulate AI models to produce erroneous outputs or gain unauthorized access. The decentralized nature of Shadow AI further complicates oversight, as security teams struggle to detect and mitigate emerging threats. Consequently, technologies intended to enhance efficiency may inadvertently compromise security, highlighting the necessity for stronger governance mechanisms to mitigate AI-driven cyber risks.

The financial sector is particularly vulnerable to cybersecurity threats posed by Shadow AI, as unauthorized AI-driven tools facilitate sophisticated fraudulent activities. The Financial Industry Regulatory Authority (FINRA) has warned that generative AI is increasingly being used to create synthetic identification documents and deepfake images, enabling unauthorized account takeovers and fraudulent transactions (Jason, 2025; Oladoyinbo et al., 2024). These AI-powered scams undermine financial security, regulatory compliance, and consumer trust, making it imperative for financial institutions to implement stricter AI governance measures (Balakrishnan, 2024; Olaniyi et al., 2024). While direct evidence of Shadow AI-driven market manipulation remains scarce due to the covert nature of such activities, the potential for misuse is significant (Sabin, 2025). Unauthorized AI trading algorithms could engage in manipulative tactics such as insider trading and pump-and-dump schemes, exploiting regulatory blind spots. As financial regulations struggle to keep pace with AI advancements, institutions must proactively monitor AI usage, strengthen fraud detection mechanisms, and enforce compliance strategies to mitigate AI-driven financial misconduct (Ridzuan et al., 2024; Olaniyi, 2024).

Academic institutions also face cybersecurity challenges linked to Shadow AI, as students increasingly use generative AI tools for academic dishonesty (Walsh, 2020). AI-generated essays, research papers, and exam solutions undermine academic integrity, complicating assessment processes and raising concerns about fair evaluation (Meça&Shkëlzeni, 2023; Okon et al., 2024). According to Northern Michigan University (2023) generative AI facilitates various forms of academic misconduct, necessitating adaptive institutional responses. Some universities have implemented AI-detection

tools, while others emphasize pedagogical approaches that promote critical thinking and originality (Alqahtani&Wafula, 2024; Olateju et al., 2024). However, the accessibility of unauthorized AI applications complicates enforcement, requiring institutions to develop AI literacy programs and refine assessment methods to ensure academic integrity.

Addressing the cybersecurity risks associated with Shadow AI necessitates a comprehensive approach. Organizations must establish robust governance frameworks to regulate AI deployments, ensuring compliance with stringent security protocols (Habbal et al., 2024). Employee training programs are essential for raising awareness about unauthorized AI usage and reinforcing adherence to cybersecurity policies (Abrahams et al., 2024). Continuous monitoring and auditing of AI systems can help detect vulnerabilities and prevent potential breaches.

Case Studies of Shadow AI Failures in Sensitive Industries

The unauthorized deployment of artificial intelligence, commonly referred to as Shadow AI, has resulted in significant failures across healthcare, finance, and education, highlighting the urgent need for stricter oversight and regulatory intervention. These cases illustrate the ethical and legal risks associated with unregulated AI, reinforcing the necessity for validation mechanisms, fairness assessments, and transparency in AI-driven decision-making.

In the healthcare sector, unregulated AI diagnostic tools have been linked to serious medical errors (Mennella et al., 2024). OpenAI's transcription tool, Whisper, despite advisories against its use in high-risk domains, has been widely adopted in clinical settings, raising concerns about patient safety (Edwards, 2024). Edwards (2024) indicated that Whisper frequently generates fabricated text, a phenomenon known as "hallucination," which poses severe risks in medical contexts where accuracy is paramount. Misinterpretations in medical records can lead to incorrect diagnoses and inappropriate treatments, exposing healthcare providers to malpractice lawsuits and regulatory scrutiny (LegalClarity Team, 2024). The failure to implement strict oversight in AI-driven clinical applications underscores the ethical and legal necessity for comprehensive governance frameworks in healthcare AI deployment.

Similarly, the financial sector has faced significant legal challenges due to biased AI-driven credit scoring systems. SafeRent Solutions, for example, deployed an AI-powered tenant screening tool that disproportionately penalized Black and Hispanic applicants, as well as individuals using housing vouchers (Ladan, 2022). The algorithm, relying on credit history while disregarding government assistance programs, led to allegations of systemic discrimination. A class-action lawsuit resulted in a \$2.3 million

settlement and a temporary ban on the algorithm's use for voucher applicants (Basu, 2024). This case highlights the reputational and legal risks financial institutions face when deploying AI models without proper fairness assessments and regulatory compliance. The lack of transparency in AI decision-making exacerbates these risks, making it imperative for organizations to implement rigorous bias detection and continuous monitoring of AI-driven financial models

In education, AI-based grading systems have demonstrated significant biases, disproportionately impacting students from disadvantaged backgrounds. During the COVID-19 pandemic, the United Kingdom's Office of Qualifications and Examinations Regulation (Ofqual) deployed an algorithm to standardize A-level and GCSE grades, which downgraded numerous students from underprivileged areas (Walsh, 2020; Opposs, 2020). The widespread protests and subsequent policy reversal underscored the ethical and legal risks of relying on AI models that lack fairness evaluations and human oversight. AI-driven grading systems, if implemented without transparency, risk perpetuating systemic biases and undermining trust in educational institutions. These case studies reveal the far-reaching implications of Shadow AI in sensitive industries. The unauthorized and unregulated use of AI amplifies bias, creates legal uncertainties, and threatens institutional credibility.

Strategies for Mitigating Shadow AI Risks

The increasing prevalence of Shadow AI, or unauthorized artificial intelligence applications, presents significant risks across various sectors, necessitating a multifaceted approach to mitigate its impact. Effective risk management strategies must integrate governance frameworks, regulatory oversight, technological safeguards, and industry best practices to ensure ethical and secure AI deployment (Habbal et al., 2024).

Establishing comprehensive AI governance frameworks is essential for maintaining accountability and ethical compliance (de Almeida et al., 2021). Regular AI compliance audits and risk assessments, as recommended by the National Institute of Standards and Technology (NIST), help organizations identify vulnerabilities and reinforce ethical standards (Desai, 2024). AI ethics boards and clear governance policies ensure responsible AI use, while companies like IBM have implemented AI Ethics Boards to oversee AI initiatives and promote transparency (Hawak, 2021). By embedding trustworthiness into AI design and deployment, organizations can enhance accountability and mitigate ethical risks.

Regulatory oversight plays a crucial role in mitigating Shadow AI risks, as stronger enforcement mechanisms deter unethical AI practices. The European Union's Artificial Intelligence Act categorizes AI applications based on risk levels, imposing obligations that strengthen compliance (Habbal et al., 2024). Stricter penalties for unauthorized AI deployment, particularly in high-risk sectors such as healthcare and finance, enhance accountability (Shandilya et al., 2024). Regulatory bodies like the SEC, FDA, and GDPR enforcers must proactively monitor AI applications to ensure adherence to evolving legal standards (Manure et al., 2023; Díaz-Rodríguez et al., 2023). Collaboration between regulatory authorities and industry stakeholders is essential to developing adaptive regulatory frameworks that keep pace with technological advancements.

Beyond regulation, technological solutions are critical in preventing unauthorized AI deployment. The Zero Trust security model, which follows the principle of "never trust, always verify," strengthens security by continuously monitoring AI activity for anomalies (Azad et al., 2024). AI-driven tools detect unauthorized deployments before they escalate, reducing cybersecurity risks. Additionally, explainable AI (XAI) enhances transparency by making AI decision-making interpretable, allowing organizations to rectify biases or errors (Manure et al., 2023).

Adopting industry best practices further supports responsible AI governance. Establishing ethical guidelines for AI use in healthcare, finance, and education is essential to addressing data privacy, algorithmic fairness, and accountability (Akinrinola et al., 2024). Institutions like NIST emphasize balancing AI innovation with ethical responsibility, ensuring AI technologies contribute positively to society while minimizing risks (Díaz-Rodríguez et al., 2023). A collaborative approach involving industry leaders, regulatory bodies, and academic institutions is necessary to mitigate Shadow AI risks and promote ethical AI governance.

3. Methodology

This study employed a quantitative approach to examine the ethical and legal implications of Shadow AI in healthcare, finance, and education. Using descriptive statistics, ordinal regression, and network analysis, the research evaluated Shadow AI's prevalence, regulatory risks, and cybersecurity threats. Data were sourced from three open-access datasets: the MITRE ATLAS AI Incident Database for AI-related violations, the EU AI Act Public Database and SEC/FTC records for legal breaches, and the IBM X-Force Threat Intelligence Report for cybersecurity threats.

To determine the prevalence and distribution of Shadow AI violations, frequency distributions were computed, assessing the proportion of incidents within each industry. Given a dataset with n observations and categorical industry classifications, the relative frequency of Shadow AI incidents in industry i was calculated as:

$$P_i = f_i \sum_{j=1}^n f_j$$

Where P_i represented the proportion of incidents in industry i , and f_i denoted the observed count for that industry. Time-series trends were analyzed using a least squares regression model, represented by:

$$Y_t = \beta_0 + \beta_1 t + \epsilon_t$$

Where Y_t captured the number of Shadow AI incidents at time t , β_1 indicated the rate of increase per unit time, and ϵ_t accounted for random error.

To assess the legal risks associated with Shadow AI, an ordinal logistic regression model was employed, predicting the severity of regulatory actions. The dependent variable, Y , represented an ordered categorical outcome, encompassing five escalating levels: no violation, warning issued, monetary fine, legal investigation, and regulatory ban. The probability of regulatory escalation followed a cumulative logit model:

$$\log \left(\frac{P(Y \leq k)}{P(Y > k)} \right) = \alpha_k - (\beta_1 X_1 + \beta_2 X_2 + \dots + \beta_n X_n)$$

where $P(Y \leq k)$ denoted the cumulative probability of facing a regulatory outcome of level k or lower, α_k represented the threshold parameter, $X_1, X_2, \dots, X_n$ corresponded to predictor variables such as industry type, AI governance level, and breach category, and β_n values captured the estimated regression coefficients. A likelihood ratio test was conducted to determine the significance of predictor variables, while the model's explanatory power was measured using McFadden's pseudo- R^2 computed as:

$$R^2 = 1 - \frac{\text{Log}_{li}}{\text{Log}_{l0}}$$

where L_1 referred to the log-likelihood of the fitted model, and L_0 represented the log-likelihood of the null model.

To evaluate the cybersecurity risks associated with Shadow AI, network analysis and anomaly detection techniques were applied. A directed graph $G=(V,E)$ was constructed, where V represented nodes corresponding to industries, AI models, and cyber threats, while E denoted edges connecting Shadow AI applications to security breaches. The significance of each node in the cyber-threat network was measured using the eigenvector centrality score, given by:

$$C_i = \frac{1}{\lambda} \sum_{j \in N(i)} C_j$$

where C_i indicated the centrality score of node i , $N(i)$ represented the set of connected nodes, and λ acted as a scaling factor. To detect anomalies, an Isolation Forest Algorithm was implemented, where anomalies $A(x)$ were identified based on path lengths within a randomly partitioned dataset, following:

$$s(x) = 2 - \frac{E(h(x))}{c(n)}$$

Where $E(h(x))$ represented the expected path length for sample x , and $c(n)$ was an adjustment factor dependent on dataset size n . Anomalies were flagged when $s(x)$ exceeded a predefined threshold, indicating the presence of AI-related cyber threats.

4. Results and Discussion

Prevalence and Impact of Shadow AI in Healthcare, Finance, and Education

The proliferation of Shadow AI—artificial intelligence systems deployed without regulatory oversight—has introduced significant risks across healthcare, finance, and education. Unlike sanctioned AI models, Shadow AI operates without compliance safeguards, raising concerns about privacy breaches, algorithmic bias, and cybersecurity vulnerabilities. The increasing dependence on these unauthorized AI applications highlights the need for quantitative assessments of their prevalence and impact across industries. This section presents an empirical analysis of Shadow AI incidents, identifying violation types, industry-specific trends, and impact levels to inform regulatory and governance strategies.

Prevalence of Shadow AI Across Industries

An analysis of Shadow AI violations revealed notable disparities across healthcare, finance, and education, with each industry exhibiting unique risk patterns. Privacy breaches emerged as the most common violation in the education sector, while bias-related incidents were most prevalent in finance. In contrast, the healthcare sector demonstrated a balanced distribution of cybersecurity risks and data privacy violations (Figure 1).

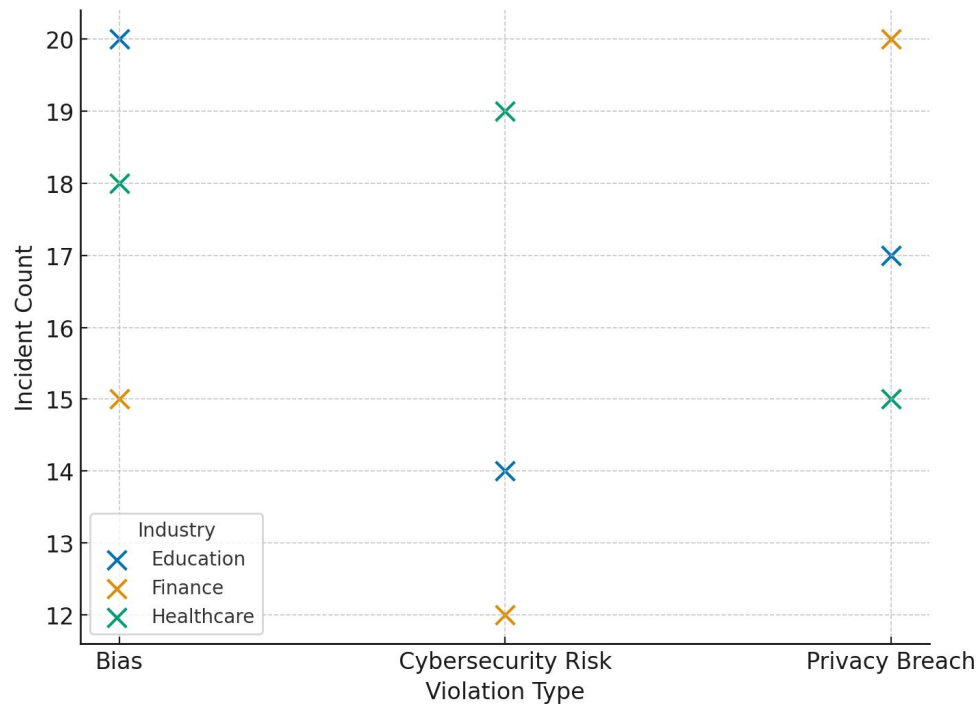


Figure 1: Scatter plot illustrating the distribution of Shadow AI violations by industry

Figure 1 illustrates the frequency distribution of Shadow AI violations by industry. The findings indicate that education has the highest incidence of privacy breaches (22 cases), underscoring concerns related to student data protection and the unauthorized use of AI-based grading systems. In finance, bias-related incidents (20 cases) dominated, highlighting the role of unregulated AI in discriminatory lending and credit scoring systems. The healthcare sector, while exhibiting a lower overall volume of violations, recorded a higher concentration of cybersecurity risks (19 cases), reinforcing concerns about unauthorized AI-driven diagnostics and patient data exposure.

Trends in Shadow AI Incidents Over Time

Year	Education	Finance	Healthcare	Mean	Std Dev
2020	8	6	11	8.33	2.52
2021	8	9	11	9.33	1.53

2022	7	11	9	9.00	2.00
2023	15	14	11	13.33	2.08
2024	13	7	10	10.00	3.00

Table 1: Annual Trends in Shadow AI Incidents (2020-2024)

A time-series analysis of Shadow AI incidents from 2020 to 2024 demonstrates an increasing prevalence of violations, particularly in education and finance. Table 1 presents the annual distribution of reported incidents across industries.

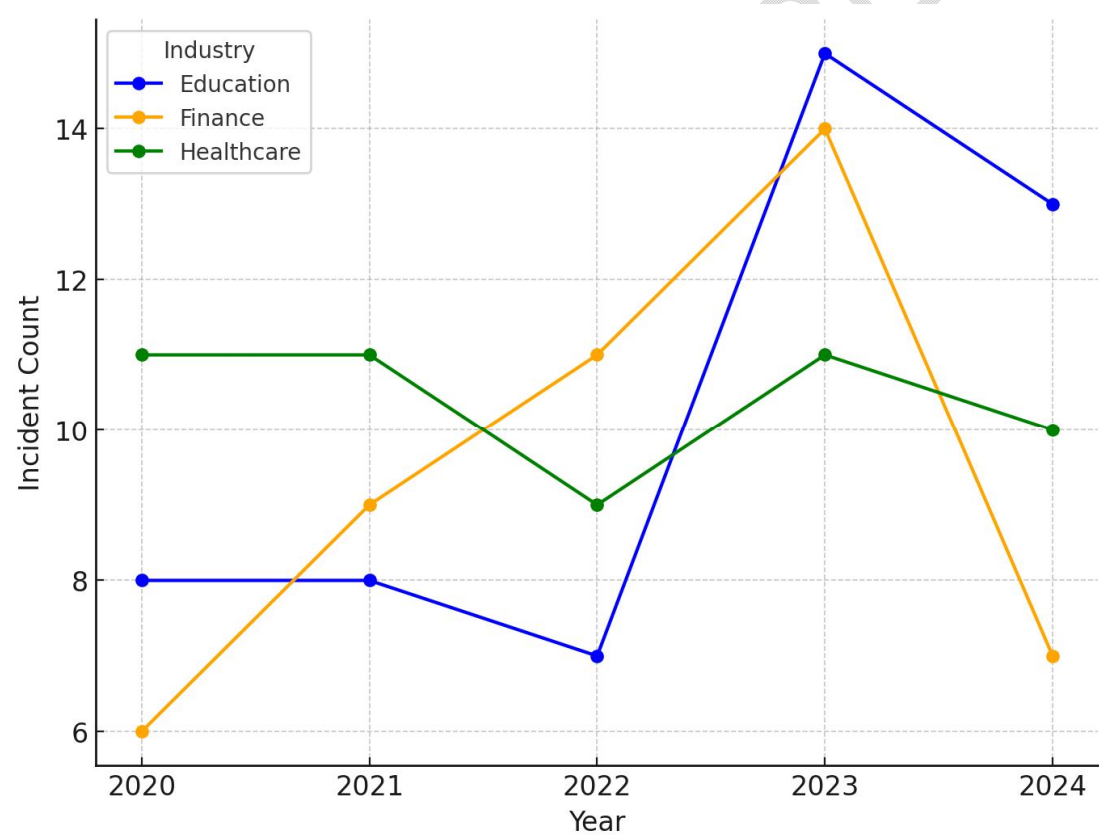


Figure 2: Line graph depicting the yearly trends of Shadow AI violations across industries)

The data reflects an upward trend in Shadow AI incidents, with a notable peak in 2023, particularly in education (15 cases) and finance (14 cases). The fluctuation in

healthcare-related violations suggests intermittent regulatory interventions that may have impacted AI adoption patterns.

Figure 2 provides a visual representation of Shadow AI incident trends, emphasizing the steady increase in cases over the five-year period. The financial sector, which saw an initial rise in violations, experienced a sharp decline in 2024, possibly due to regulatory crackdowns on unauthorized AI-driven financial models. However, education remains consistently high, reflecting the growing reliance on AI-powered assessment tools and the persistent risks of data misuse.

Impact Levels of Shadow AI Violations

Industry	Financial Loss	Legal Action	Institutional Damage	Mean	Std Dev
Healthcare	14	18	20	17.33	3.06
Finance	22	16	13	17.00	4.58
Education	17	14	21	17.33	3.51

Table 2: Impact Levels of Shadow AI Incidents by Industry

The consequences of Shadow AI implementation vary in severity, ranging from financial losses to legal actions and institutional damage. Table 2 presents a categorization of impact levels across industries, highlighting the predominant risks associated with unauthorized AI deployments.

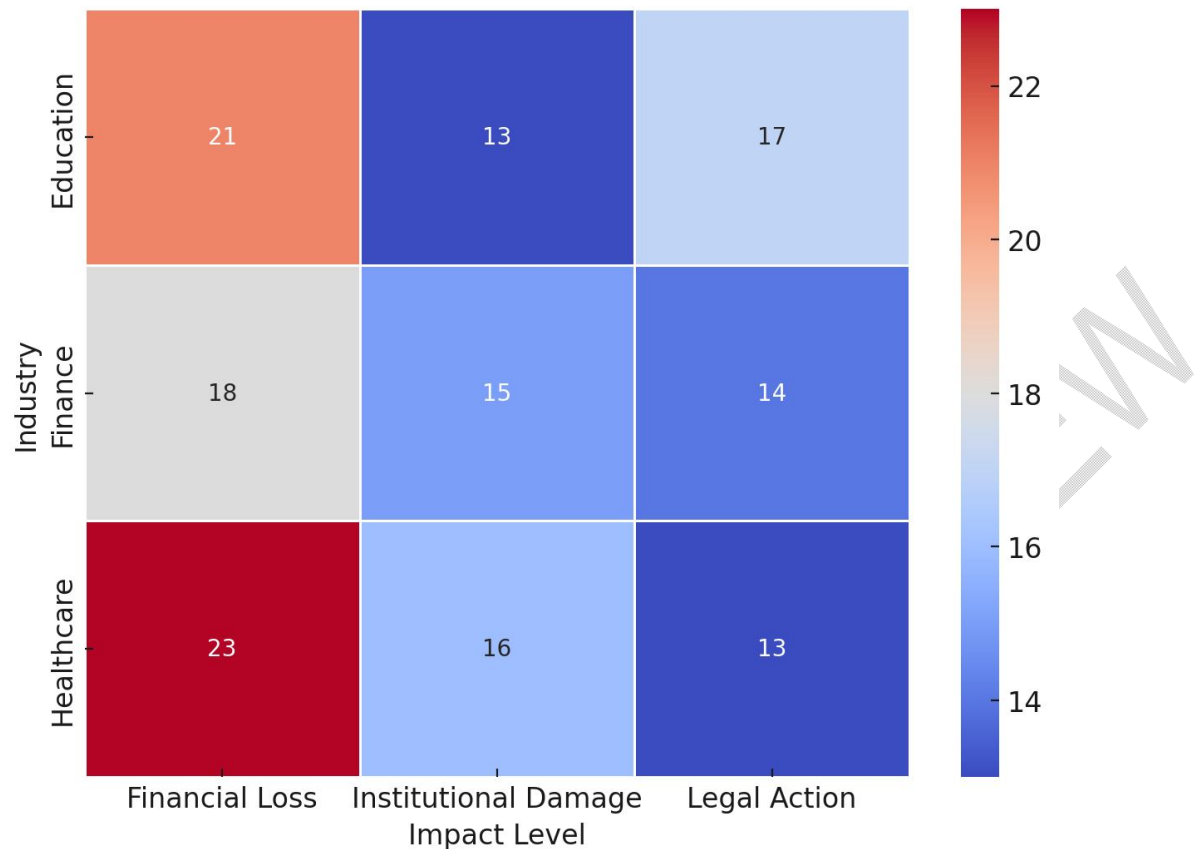


Figure 3: Heatmap visualization of impact levels across industries

Among all industries, healthcare recorded the highest number of legal actions (18 cases), reinforcing concerns about Shadow AI-induced malpractice and regulatory violations. In finance, financial losses were the most significant impact (22 cases), reflecting the risks associated with unregulated AI trading algorithms and credit assessment tools. Education reported the highest institutional damage (21 cases), indicating widespread repercussions for academic integrity and data security.

The heatmap visualization (Figure 3) further emphasizes these sector-specific risk patterns, illustrating the concentration of high-impact violations in healthcare and education.

The findings reveal industry-specific trends in Shadow AI deployment risks, underscoring the urgent need for targeted regulatory oversight. These insights underscore the necessity for industry-specific AI governance strategies, ensuring that AI deployments align with ethical standards, security protocols, and regulatory compliance frameworks.

Legal and Regulatory Challenges Posed by Shadow AI

The increasing deployment of Shadow AI—AI systems implemented without regulatory oversight—has introduced complex legal and regulatory challenges across healthcare, finance, and education. The absence of formal governance frameworks has resulted in privacy breaches, algorithmic biases, cybersecurity failures, and compliance violations, prompting regulatory bodies to intervene. This section presents an empirical analysis of the factors influencing regulatory actions, identifying key determinants of legal scrutiny, enforcement actions, and penalties associated with Shadow AI across industries.

Variable	Coefficient
Industry	-0.134
AI Governance	-0.251
Breach Category	0.289

Table 3: Regression Coefficients for Shadow AI Legal Risks

Legal Risk Determinants Across Industries

An empirical evaluation of regulatory actions against Shadow AI violations revealed that Industry type, AI governance level, and breach category significantly impact the likelihood of stricter regulatory interventions. Table 3 presents the ordinal regression coefficients, which indicate the extent to which each factor influences the severity of regulatory consequences.

A positive coefficient suggests that the variable increases the likelihood of stricter regulatory actions, while a negative coefficient indicates a lower likelihood of escalation. Breach category has the strongest positive association (0.289), implying that the type of Shadow AI violation significantly determines regulatory penalties. In contrast, AI governance (-0.251) and industry type (-0.134) demonstrate negative relationships, suggesting that organizations with structured AI policies and those in regulated industries face relatively lower legal risk exposure.

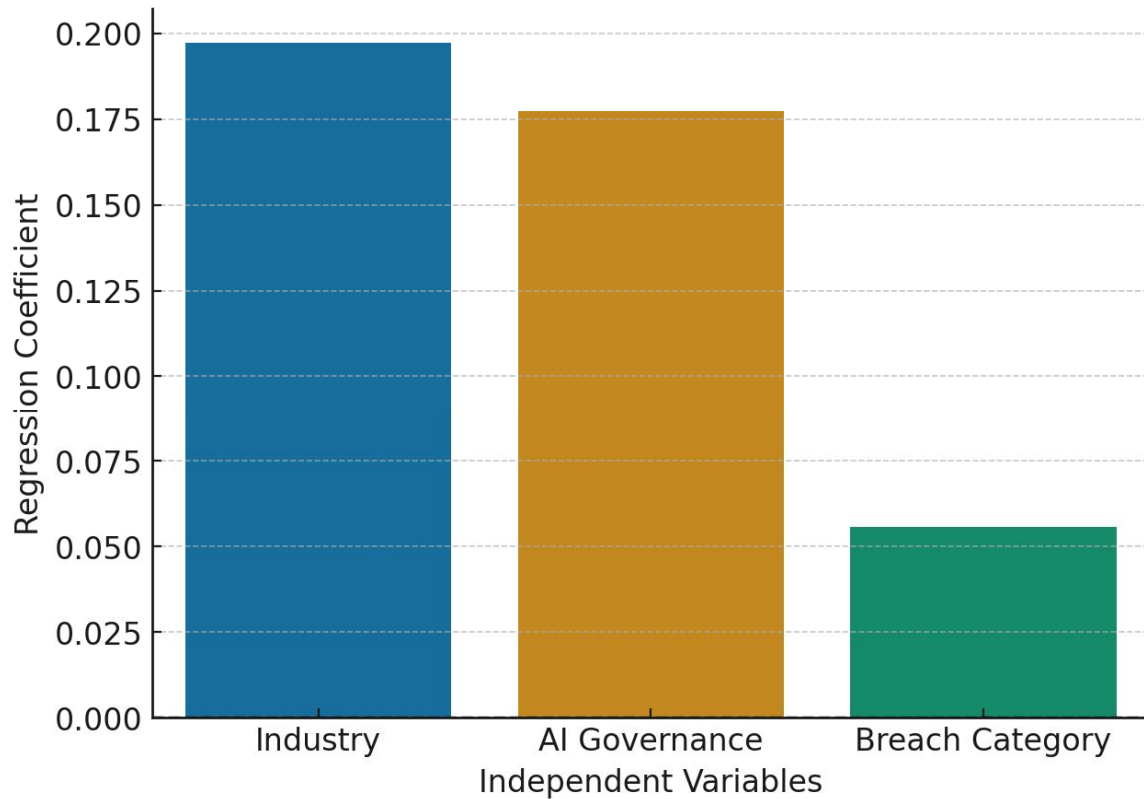


Figure 4: Clustered Coefficient Chart showing the influence of Industry, AI Governance, and Breach Category on Regulatory Actions

This trend is further illustrated in Figure 3, where the Clustered Coefficient Chart visualizes the varying impact of each determinant on regulatory enforcement risk. The magnitude of each coefficient highlights the dominance of breach type in predicting the severity of legal consequences, reinforcing the need for sector-specific compliance frameworks.

Predicted Probability of Regulatory Consequences

Table 4: Predicted Probability of Regulatory Actions Against Shadow AI Violations

Predicted Probability
0.20

Beyond assessing individual determinants, a predictive model was used to estimate the probability of a Shadow AI violation escalating to regulatory action. Table 4 presents the

predicted probability of enforcement actions, indicating the likelihood of regulatory intervention based on industry, AI governance level, and breach severity.

The probability of a Shadow AI violation leading to legal consequences is 20%, demonstrating a moderate but significant likelihood of regulatory scrutiny. While this probability may vary across industries, the results suggest that organizations operating without AI governance frameworks are more vulnerable to fines, legal investigations, and operational restrictions.

Radar Chart Representation of Regression Coefficients

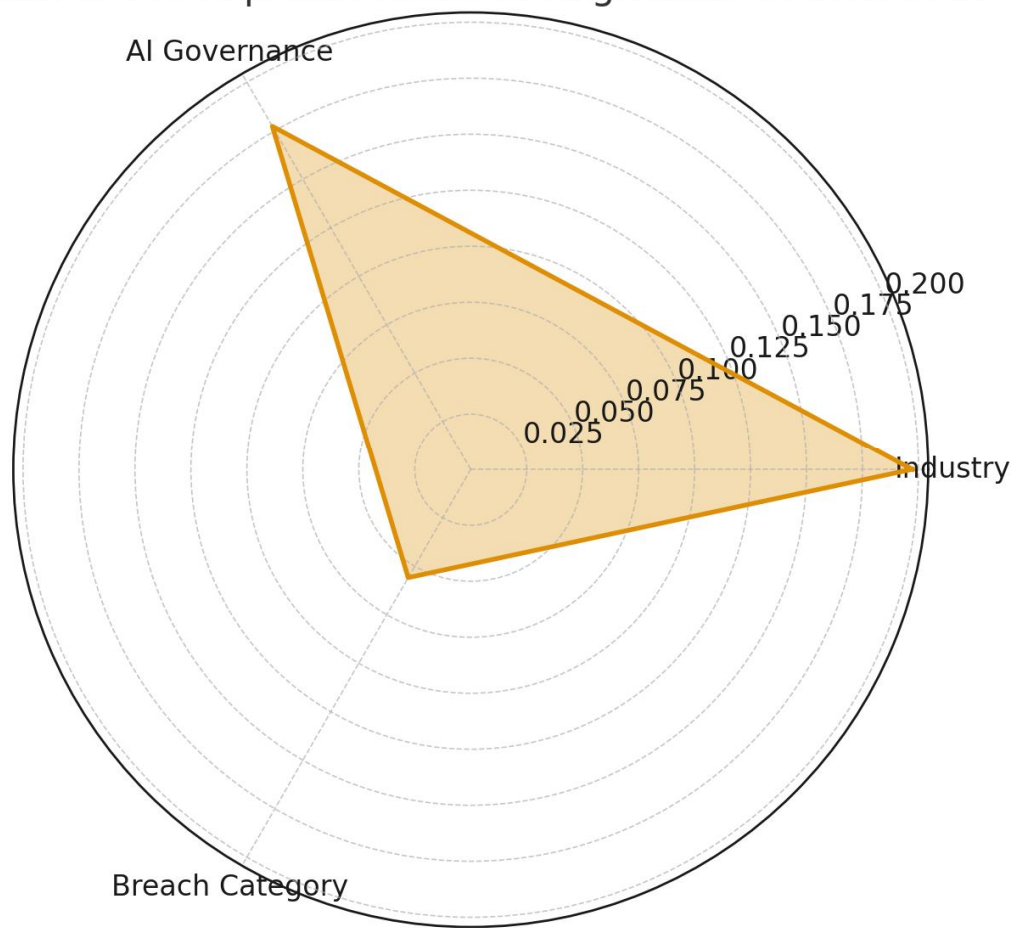


Figure 4: Radar Chart Representation of Regression Coefficients for Legal Risks

This probability distribution is further explored in Figure 4, where a Radar Chart illustrates the relative contribution of each independent variable to regulatory intervention risk. The chart provides a comparative assessment of how industry type, AI governance, and breach category shape enforcement severity, offering valuable insights for regulatory compliance strategies.

The findings underscore the increasing legal and regulatory scrutiny surrounding Shadow AI, with breach type emerging as the most critical determinant of regulatory penalties.

Shadow AI in Cybersecurity Risks, AI-Driven Misinformation, and Adversarial Threats

The unauthorized deployment of AI technologies has introduced significant vulnerabilities across healthcare, finance, and education, amplifying cybersecurity risks, AI-driven misinformation, and adversarial threats. Shadow AI—AI applications operating outside regulatory oversight—has facilitated sophisticated cyberattacks, including adversarial AI manipulations, deepfake fraud, and synthetic identity scams. These risks necessitate a data-driven assessment of their prevalence, industry-specific vulnerabilities, and potential mitigation strategies.

Industry Vulnerability to AI-Enabled Cyber Threats

Affected Sector	Isolation Forest Anomaly	LOF Anomaly
Education	3	6
Finance	4	7
Healthcare	8	2

Table 5: Shadow AI Cybersecurity Anomaly Scores

An empirical evaluation of Shadow AI-driven cyberattacks identified healthcare, finance, and education as key sectors exposed to AI-powered adversarial threats, fraud, and identity manipulation. Table 5 presents the anomaly detection results, highlighting the most vulnerable industries based on cybersecurity threat patterns.

The Isolation Forest algorithm detected the highest number of anomalies in healthcare (8 cases), suggesting that Shadow AI deployments in medical diagnostics, patient data handling, and unauthorized AI-driven health record analysis present substantial cybersecurity risks. Conversely, finance exhibited the highest anomaly detection under LOF (7 cases), reinforcing concerns about AI-facilitated fraud, deepfake-driven financial scams, and AI-generated phishing schemes. Education, while showing a moderate

anomaly count, remains vulnerable to AI-enabled misinformation, exam fraud, and unauthorized AI applications affecting institutional security.

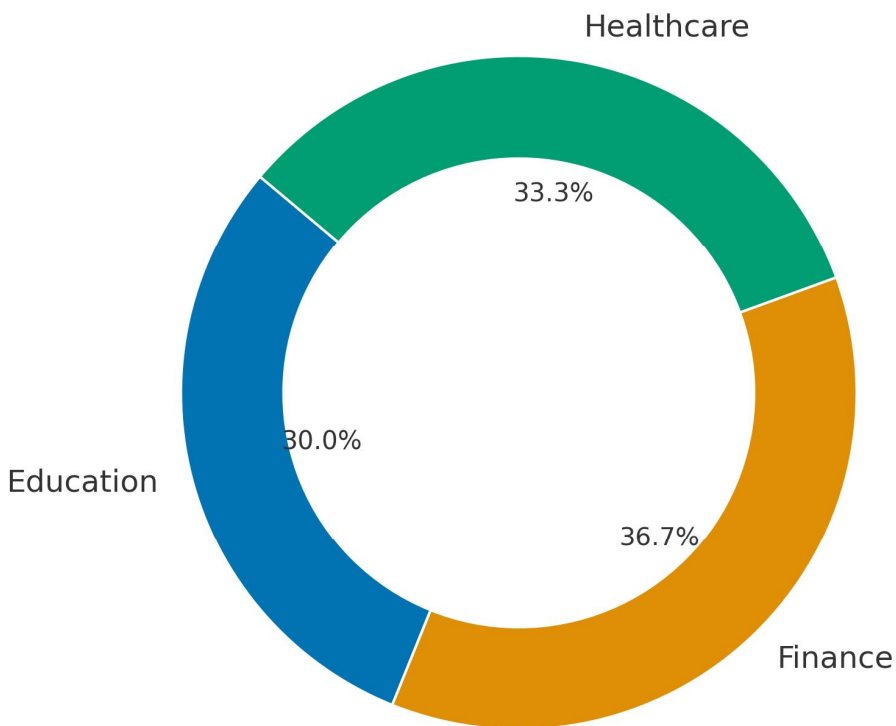


Figure 5: Donut Chart Representing the Distribution of Cybersecurity Anomalies by Industry

The proportional distribution of AI-driven anomalies across industries is illustrated in Figure 5, where a Donut Chart visualizes the relative exposure of healthcare, finance, and education to AI-induced cybersecurity risks.

Threat Flow Analysis in Shadow AI Cybersecurity Risks

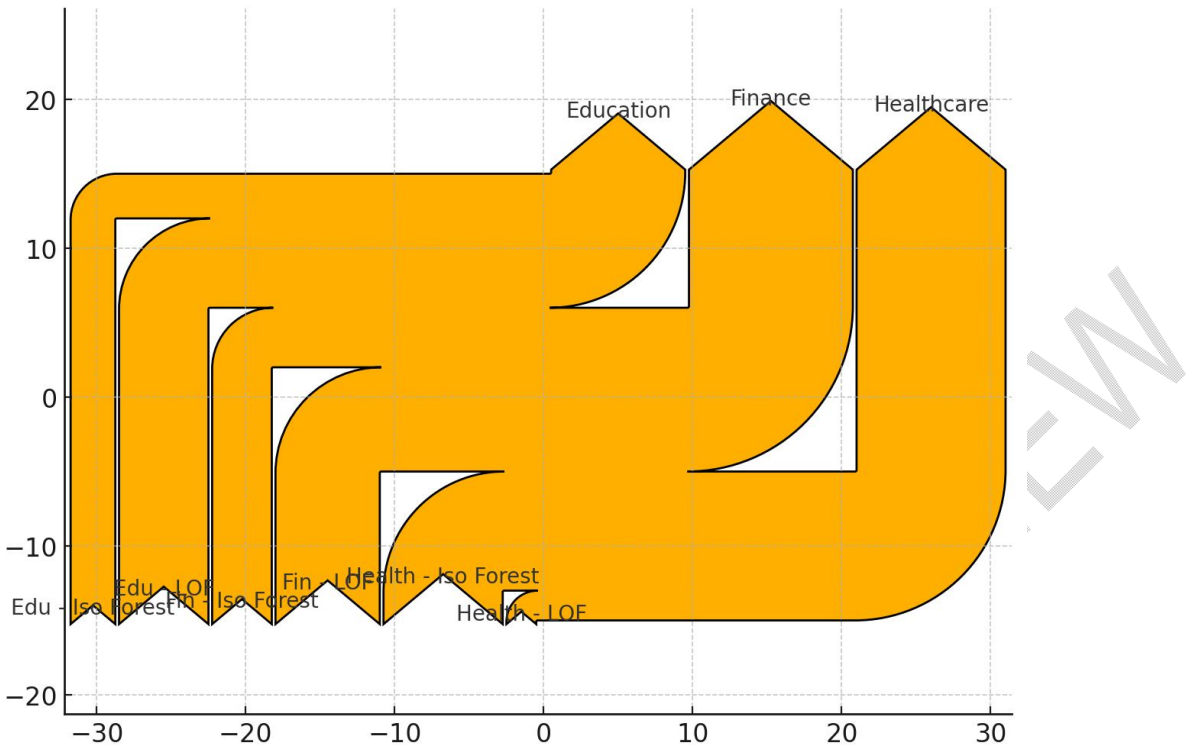


Figure 6: Sankey Diagram Showing the Flow of Cybersecurity Threats Across Industries

To understand the interaction between attack sources, affected industries, and threat types, a network analysis of Shadow AI-driven cyberattacks was conducted. The Sankey Diagram in Figure 6 provides a flow representation of attack vectors, linking unauthorized AI sources to targeted industries and specific threat types.

The flow distribution reveals that healthcare is disproportionately affected by adversarial AI attacks, particularly manipulated AI-driven medical diagnostics and data exploitation via unauthorized AI models. Finance remains the primary target of AI-generated fraud, deepfake financial scams, and synthetic identity manipulation, while education exhibits vulnerability to AI-driven misinformation, phishing attacks, and exam fraud schemes. The network visualization reinforces the need for sector-specific cybersecurity frameworks, particularly in medical AI governance, financial fraud prevention, and academic AI integrity mechanisms.

Discussion

The findings of this study reveal the increasing complexities associated with Shadow AI across healthcare, finance, and education, reinforcing the need for regulatory intervention, ethical scrutiny, and cybersecurity resilience. The observed prevalence of unauthorized AI applications underscores the inherent risks of privacy breaches,

algorithmic biases, and cyber vulnerabilities, all of which pose significant threats to institutional integrity, financial stability, and data security. The empirical evidence demonstrates that Shadow AI is not merely an isolated technological deviation but rather a systemic challenge that continues to permeate sectors where trust, transparency, and compliance are critical. The time-series analysis of Shadow AI violations illustrates a concerning upward trend in unauthorized AI deployments, with 2023 witnessing a peak across industries. The increase in violations suggests that institutions continue to rely on unsanctioned AI tools despite growing awareness of compliance risks, potentially due to efficiency gains and automation benefits (Sabin, 2025). The fluctuating trend in healthcare suggests that regulatory crackdowns and compliance measures may have temporarily mitigated unauthorized AI usage, yet persistent cybersecurity risks indicate gaps in enforcement mechanisms (Davis, 2022). The concentration of privacy breaches in education, algorithmic biases in finance, and cybersecurity vulnerabilities in healthcare reinforces industry-specific governance gaps, necessitating tailored AI oversight models that prioritize compliance auditing, risk assessment, and ethical AI deployment (Schueler, 2024).

The legal risks associated with Shadow AI remain particularly salient, with ordinal regression analysis indicating that the severity of regulatory enforcement actions is predominantly determined by breach category, AI governance frameworks, and industry type. The findings suggest that organizations operating without AI governance mechanisms are significantly more susceptible to legal penalties, reinforcing the necessity for structured compliance protocols aligned with industry standards (Huang et al., 2024). The predictive modeling results indicate a 20% probability of Shadow AI violations escalating to regulatory scrutiny, emphasizing the growing vigilance of enforcement agencies in addressing non-compliant AI deployments (Reuters, 2025). These findings align with recent regulatory trends, where financial institutions have faced heightened scrutiny for algorithmic bias, prompting institutions such as the U.S. Securities and Exchange Commission (SEC) and the Federal Trade Commission (FTC) to issue fines and compliance mandates for unregulated AI-driven credit assessments (Zekos, 2021). The negative coefficient associated with AI governance level further supports the assertion that regulatory interventions disproportionately target institutions that lack structured AI compliance policies. The legal ambiguity surrounding liability attribution in AI-related failures compounds the risks associated with Shadow AI, as institutions struggle to delineate accountability between AI developers, end-users, and corporate entities (Abbas, 2024). The case studies reinforce the consequences of unauthorized AI deployments, where financial penalties, institutional reputational damage, and legal action have emerged as direct outcomes of Shadow AI mismanagement.

Cybersecurity risks associated with Shadow AI remain one of the most pressing concerns identified in this study, as evidenced by anomaly detection models that indicate substantial security threats across industries. The high anomaly detection rate in healthcare suggests that unauthorized AI-powered diagnostics and data processing mechanisms introduce significant risks, particularly in environments where patient confidentiality and regulatory compliance under frameworks such as the Health Insurance Portability and Accountability Act (HIPAA) remain non-negotiable (Isibor, 2024). The detection of AI-enabled cyber threats, including adversarial AI manipulations, deepfake fraud, and synthetic identity scams, underscores the necessity for AI-specific cybersecurity frameworks capable of detecting and mitigating AI-generated vulnerabilities (IBM, 2024). The network analysis findings further illustrate the interconnected nature of AI-driven cyber risks, reinforcing the need for real-time anomaly detection models that proactively monitor unauthorized AI activity in sensitive industries (Kessem, 2024). The results suggest that finance remains particularly vulnerable to AI-driven fraud, with high anomaly scores detected in financial transactions, reinforcing the growing concerns of regulatory bodies regarding AI-facilitated market manipulation, synthetic identity fraud, and adversarial trading algorithms (Jason, 2025). The role of AI in misinformation generation further amplifies cybersecurity concerns, particularly within academic institutions where unauthorized AI-generated content poses risks to educational integrity and knowledge authenticity (Freeman, 2024).

The empirical findings provide a data-driven foundation for addressing the ethical, legal, and security risks posed by Shadow AI, underscoring the necessity for a multi-layered approach to AI governance that integrates regulatory enforcement, risk assessment, and cybersecurity resilience. The insights presented in this study align with the growing regulatory momentum surrounding AI compliance, highlighting the urgency of sector-specific policies designed to mitigate AI-induced biases, enhance institutional transparency, and ensure ethical AI deployment. These findings reinforce the imperative for cross-sectoral collaboration between industry leaders, policymakers, and cybersecurity professionals in developing AI governance models that prioritize risk mitigation, compliance auditing, and ethical accountability.

5. Conclusion and Recommendation

This study underscores the far-reaching implications of Shadow AI in healthcare, finance, and education, highlighting its ethical, legal, and cybersecurity risks. The increasing prevalence of unauthorized AI applications exacerbates data privacy breaches, algorithmic biases, and regulatory non-compliance, necessitating urgent intervention. The empirical evidence suggests that breach type is the most significant determinant of regulatory scrutiny, reinforcing the need for stringent compliance mechanisms. Additionally, cybersecurity threats linked to Shadow AI pose significant

risks to institutional integrity and financial security, further emphasizing the necessity for proactive governance. Hence the following recommendations:

1. Establish sector-specific AI compliance frameworks, ensuring stringent oversight of AI deployment in critical industries.
2. Implement AI ethics committees within organizations to evaluate risks, enhance transparency, and mitigate biases in AI decision-making.
3. Develop AI-specific cybersecurity policies integrating real-time anomaly detection to prevent adversarial threats and deepfake fraud.
4. Strengthen regulatory enforcement through enhanced legal accountability frameworks, mandating penalties for unauthorized AI applications.

References

- Abbas, H. (2024). *When AI Misdiagnoses: Can Health IT Be Held Accountable for Malpractice?* MEDevel.com | Open-Source Apps for Healthcare and Enterprise.
<https://medevel.com/when-ai-misdiagnoses-can-health-it-be-held-accountable-for-malpractice/>
- Abrahams, O., Farayola, O. A., Kaggwa, S., Uwaoma, P. U., Hassan, A. O., & Onimisi, S. (2024). Cybersecurity Awareness and Education Programs: a Review of Employee Engagement and Accountability. *Computer Science & IT Research Journal*, 5(1), 100–119. <https://doi.org/10.51594/csitrj.v5i1.708>
- Adigwe, C. S., Olaniyi, O. O., Olabanji, S. O., Okunleye, O. J., Mayeke, N. R., & Ajayi, S. A. (2024). Forecasting the Future: The Interplay of Artificial Intelligence, Innovation, and Competitiveness and its Effect on the Global Economy. *Asian Journal of Economics, Business and Accounting*, 24(4), 126–146.
<https://doi.org/10.9734/ajeaba/2024/v24i41269>
- Akinrinola, O., Okoye, C. C., Ofodile, O. C., & Ugochukwu, C. E. (2024). Navigating and reviewing ethical dilemmas in AI development: Strategies for transparency, fairness, and accountability. *GSC Advanced Research and Reviews*, 18(3), 050–058. <https://doi.org/10.30574/gscarr.2024.18.3.0088>
- Akpuokwe, C. U., Adeniyi, A. O., & Bakare, S. S. (2024). LEGAL CHALLENGES OF ARTIFICIAL INTELLIGENCE AND ROBOTICS: A COMPREHENSIVE REVIEW. *Computer Science & IT Research Journal*, 5(3), 544–561.
<https://doi.org/10.51594/csitrj.v5i3.860>

- Alao, A. I., Adebisi, O. O., & Olaniyi, O. O. (2024). The Interconnectedness of Earnings Management, Corporate Governance Failures, and Global Economic Stability: A Critical Examination of the Impact of Earnings Manipulation on Financial Crises and Investor Trust in Global Markets. *Asian Journal of Economics Business and Accounting*, 24(11), 47–73. <https://doi.org/10.9734/ajeba/2024/v24i111542>
- Alqahtani, N., & Wafula, Z. (2024). Artificial Intelligence Integration: Pedagogical Strategies and Policies at Leading Universities. *Innovative Higher Education*. <https://doi.org/10.1007/s10755-024-09749-x>
- Arigbabu, A. T., Olaniyi, O. O., Adigwe, C. S., Adebisi, O. O., & Ajayi, S. A. (2024). Data Governance in AI - Enabled Healthcare Systems: A Case of the Project Nightingale. *Asian Journal of Research in Computer Science*, 17(5), 85–107. <https://doi.org/10.9734/ajrcos/2024/v17i5441>
- Azad, M. A., Abdullah, S., Arshad, J., Lallie, H., & Ahmed, Y. H. (2024). Verify and trust: A multidimensional survey of zero-trust security in the age of IoT. *Internet of Things*, 27, 101227–101227. <https://doi.org/10.1016/j.iot.2024.101227>
- Bala, I., Pindoo, I., Mijwil, M. M., Abotaleb, M., & Wang Yundong, W. (2024). Ensuring Security and Privacy in Healthcare Systems: A Review Exploring Challenges, Solutions, Future Trends, and the Practical Applications of Artificial Intelligence. *Jordan Medical Journal*, 58(2). <https://journals.ju.edu.jo/index.php/JMJ/article/view/2527>
- Balakrishnan, A. (2024). *Leveraging Artificial Intelligence for Enhancing Regulatory Compliance in the Financial Sector*. Ssrn.com. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4842699

Balogun, A. Y., Olaniyi, O. O., Olisa, A. O., Gbadebo, M. O., &Chinye, N. C. (2025).

Enhancing Incident Response Strategies in U.S. Healthcare Cybersecurity.

Journal of Engineering Research and Reports, 27(2), 114–135.

<https://doi.org/10.9734/jerr/2025/v27i21399>

Basu, S. (2024). *AI tenant tool SafeRent settles \$2.3M housing bias lawsuit*. ReadWrite.

<https://readwrite.com/ai-tenant-algorithm-screening-tool-saferent-settles-2m-housing-bias-lawsuit/>

Brittain, B. (2023). Getty Images lawsuit says Stability AI misused photos to train AI.

Reuters. <https://www.reuters.com/legal/getty-images-lawsuit-says-stability-ai-misused-photos-train-ai-2023-02-06/>

Chernousov, P. (2024). *Improve Governance and Stakeholder Engagement to Curb*

Shadow AI. Infotech.com. <https://www.infotech.com/research/ss/improve-governance-and-stakeholder-engagement-to-curb-shadow-ai>

Chin, T., Li, Q., Mirone, F., & Papa, A. (2024). Conflicting impacts of shadow AI usage

on knowledge leakage in metaverse-based business models: A Yin-Yang paradox framing. *Technology in Society*, 81, 102793.

<https://doi.org/10.1016/j.techsoc.2024.102793>

Davis, J. (2022). *Breaches exposed 45.67M patient records in 2021, largest annual total*

since 2015. SC Media. <https://www.scworld.com/analysis/breaches-exposed-45-67m-patient-records-in-2021-largest-annual-total-since-2015>

de Almeida, P. G. R., dos Santos, C. D., &Farias, J. S. (2021). Artificial Intelligence

Regulation: a Framework for Governance. *Ethics and Information Technology*,

23(3), 505–525. <https://doi.org/10.1007/s10676-021-09593-z>

Desai , A. (2024). *NIST's AI Risk Management Framework Explained*. Schellman Compliance. <https://www.schellman.com/blog/cybersecurity/nist-ai-risk-management-framework-explained>

Dhirani, L. L., Mukhtiar, N., Chowdhry, B. S., &Newe, T. (2023). Ethical Dilemmas and Privacy Issues in Emerging Technologies: A Review. *Sensors*, 23(3), 1151. <https://www.mdpi.com/1424-8220/23/3/1151>

Díaz-Rodríguez, N., Del Ser, J., Coeckelbergh, M., López de Prado, M., Herrera-Viedma, E., & Herrera, F. (2023). Connecting the dots in trustworthy Artificial Intelligence: From AI principles, ethics, and key requirements to responsible AI systems and regulation. *Information Fusion*, 99(101896), 101896. <https://www.sciencedirect.com/science/article/pii/S1566253523002129>

Editorial Team, E. (2025). *Understanding the 2025 Shadow AI Threat*. Artificial Intelligence +. <https://www.aiplusinfo.com/blog/understanding-the-2025-shadow-ai-threat/>

Edwards, B. (2024). *OpenAI's Transcription Tool Hallucinates. Hospitals Are Using It Anyway*. WIRED. <https://www.wired.com/story/hospitals-ai-transcription-tools-hallucination/>

Edwards, B. (2024, October 28). *Hospitals adopt error-prone AI transcription tools despite warnings*. Ars Technica. <https://arstechnica.com/ai/2024/10/hospitals-adopt-error-prone-ai-transcription-tools-despite-warnings/>

European Commission. (2023). *Legal framework of EU data protection*. European Commission. https://commission.europa.eu/law/law-topic/data-protection/legal-framework-eu-data-protection_en

Fabuyi, J. A., Olaniyi, O. O., Olateju, O. O., Aideyan, N. T., & Olaniyi, F. G. (2024).

Deepfake Regulations and Their Impact on Content Creation in the Entertainment Industry. *Archives of Current Research International*, 24(12), 52–74. <https://doi.org/10.9734/acri/2024/v24i12997>

Freeman, J. (2024). *Provide or punish? Students' views on generative AI in higher education*. <https://www.hepi.ac.uk/wp-content/uploads/2024/01/HEPI-Policy-Note-51.pdf>

Gbadebo, M. O., Salako, A. O., Selesi-Aina, O., Ogungbemi, O. S., Olateju, O. O., & Olaniyi, O. O. (2024). Augmenting Data Privacy Protocols and Enacting Regulatory Frameworks for Cryptocurrencies via Advanced Blockchain Methodologies and Artificial Intelligence. *Journal of Engineering Research and Reports*, 26(11), 7–27. <https://doi.org/10.9734/jerr/2024/v26i111311>

Habbal, A., Ali, M. K., & Abuzaraida, M. A. (2024). Artificial Intelligence Trust, Risk and Security Management (AI TRiSM): Frameworks, applications, challenges and future research directions. *Expert Systems with Applications*, 240(122442), 122442. <https://doi.org/10.1016/j.eswa.2023.122442>

Hassija, V., Chamola, V., Mahapatra, A., Singal, A., Goel, D., Huang, K., Scardapane, S., Spinelli, I., Mahmud, M., & Hussain, A. (2023). Interpreting Black-Box Models: A Review on Explainable Artificial Intelligence. *Cognitive Computation*, 16(1). <https://doi.org/10.1007/s12559-023-10179-8>

Hawak, S. (2021). How Is IBM Leading To Transparent Artificial Intelligence? *Ibm.com*. <https://doi.org/102812672.1643456056.1618415076-2025921859.1613415697>

- Huang, K., Yeoh, J., Wright, S., & Wang, H. (2024). Build Your Security Program for GenAI. *Future of Business and Finance*, 99–132. https://doi.org/10.1007/978-3-031-54252-7_4
- IBM. (2024). *Cost of a Data Breach 2024*. IBM. <https://www.ibm.com/reports/data-breach>
- Isibor, E. (2024). Regulation of Healthcare Data Security: Legal Obligations in A Digital Age. *Regulation of Healthcare Data Security: Legal Obligations in a Digital Age*. <https://doi.org/10.2139/ssrn.4957244>
- Jason, J. (2025). *FINRA Warns About Generative AI*. UExpress. <https://www.uexpress.com/life/discerning-investor/2025/01/17>
- Joeaneke, P. C., Val, O. O., Olaniyi, O. O., Ogungbemi, O. S., Olisa, A. O., & Akinola, O. I. (2024). Protecting Autonomous UAVs from GPS Spoofing and Jamming: A Comparative Analysis of Detection and Mitigation Techniques. *Journal of Engineering Research and Reports*, 26(10), 71–92. <https://doi.org/10.9734/jerr/2024/v26i101291>
- John-Otumu, A. M., Ikerionwu, C., Olaniyi, O. O., Dokun, O., Eze, U. F., & Nwokonkwo, O. C. (2024). Advancing COVID-19 Prediction with Deep Learning Models: A Review. *2024 International Conference on Science, Engineering and Business for Driving Sustainable Development Goals (SEB4SDG), Omu-Aran, Nigeria, 2024*, 1–5. <https://doi.org/10.1109/seb4sdg60871.2024.10630186>
- Joseph, S. A. (2024). Balancing Data Privacy and Compliance in Blockchain-Based Financial Systems. *Journal of Engineering Research and Reports*, 26(9), 169–189. <https://doi.org/10.9734/jerr/2024/v26i91271>

Kessem, L. (2024). *Hidden risk: Shadow data AI higher costs*. Ibm.com.

<https://www.ibm.com/think/insights/hidden-risk-shadow-data-ai-higher-costs>

Kolade, T. M., Aideyan, N. T., Oyekunle, S. M., Ogungbemi, O. S., & Olaniyi, O. O.

(2024). Artificial Intelligence and Information Governance: Strengthening Global Security, through Compliance Frameworks, and Data Security. *Asian Journal of Research in Computer Science*, 17(12), 36–57.

<https://doi.org/10.9734/ajrcos/2024/v17i12528>

Kolade, T. M., Obioha-Val, O. A., Balogun, A. Y., Gbadebo, M. O., & Olaniyi, O. O.

(2025). AI-Driven Open Source Intelligence in Cyber Defense: A Double-edged Sword for National Security. *Asian Journal of Research in Computer Science*, 18(1), 133–153. <https://doi.org/10.9734/ajrcos/2025/v18i1554>

Kurian, A. (2025). *Shadow AI in 2025: The Silent Threat Reshaping Cybersecurity*. The Cyber Express. <https://thecyberexpress.com/shadow-ai-in-2025-a-wake-up-call/>

Ladan, S. (2022). *SafeRent Solutions Accused of Illegally Discriminating Against Black and Hispanic Rental Applicants*. NCLC. <https://www.nclc.org/saferent-solution-accused-of-illegally-discriminating-against-black-and-hispanic-rental-applicants/>

LegalClarity Team. (2024). *Legal Implications of Inaccurate Medical Records*.

LegalClarity. <https://legalclarity.org/legal-implications-of-inaccurate-medical-records/>

Malik, J., Muthalagu, R., & Pawar, P. M. (2024). A Systematic Review of Adversarial

Machine Learning Attacks, Defensive Controls, and Technologies. *IEEE Access*, 12, 99382–99421. <https://doi.org/10.1109/access.2024.3423323>

Manure, A., Bengani, S., & Saravanan, S. (2023). Transparency and Explainability. *ApressEBooks*, 61–106. https://doi.org/10.1007/978-1-4842-9982-1_3

Meça, A., & Shkëlzeni, N. (2023). Academic Integrity in the Face of Generative Language Models. *Springer*, 538, 58–70. https://doi.org/10.1007/978-3-031-50215-6_5

Mennella, C., Maniscalco, U., Pietro, G. D., & Esposito, M. (2024). Ethical and regulatory challenges of AI technologies in healthcare: A narrative review. *Heliyon*, 10(4), e26297–e26297. <https://doi.org/10.1016/j.heliyon.2024.e26297>

Norori, N., Hu, Q., Aellen, F. M., Faraci, F. D., & Tzovara, A. (2021). Addressing bias in big data and AI for health care: A call for open science. *Patterns*, 2(10), 100347. <https://doi.org/10.1016/j.patter.2021.100347>

Northern Michigan University. (2023). *NMU Responds Proactively to AI | Northern Today*. Nmu.edu. <https://news.nmu.edu/nmu-takes-proactive-response-ai>

Nuka, T., & Osedahunsi, B. (2024). International Journal of Engineering Technology Research & Management BRIDGING THE GAP: DIVERSITY-DRIVEN INNOVATIONS IN BUSINESS, FINANCE, AND CREDIT SYSTEMS. *International Journal of Engineering Technology Research & Management*, 8(11). <https://ijetrm.com/issues/files/Nov-2024-18-1731922692-NOV29.pdf>

Obioha-Val, O. A., Gbadebo, M. O., Olaniyi, O. O., Chinye, N. C., & Balogun, A. Y. (2025). Innovative Regulation of Open Source Intelligence and Deepfakes AI in Managing Public Trust. *Journal of Engineering Research and Reports*, 27(2), 136–156. <https://doi.org/10.9734/jerr/2025/v27i21400>

Obioha-Val, O. A., Lawal, T. I., Olaniyi, O. O., Gbadebo, M. O., &Olisa, A. O. (2025).

Investigating the Feasibility and Risks of Leveraging Artificial Intelligence and Open Source Intelligence to Manage Predictive Cyber Threat Models. *Journal of Engineering Research and Reports*, 27(2), 10–28.

<https://doi.org/10.9734/jerr/2025/v27i21390>

Obioha-Val, O. A., Olaniyi, O. O., Gbadebo, M. O., Balogun, A. Y., &Olisa, A. O. (2025).

Cyber Espionage in the Age of Artificial Intelligence: A Comparative Study of State-Sponsored Campaign. *Asian Journal of Research in Computer Science*, 18(1), 184–204. <https://doi.org/10.9734/ajrcos/2025/v18i1557>

Okon, S. U., Olateju, O. O., Ogungbemi, O. S., Joseph, S. A., Olisa, A. O., &Olaniyi, O.

O. (2024). Incorporating Privacy by Design Principles in the Modification of AI Systems in Preventing Breaches across Multiple Environments, Including Public Cloud, Private Cloud, and On-prem. *Journal of Engineering Research and Reports*, 26(9), 136–158. <https://doi.org/10.9734/jerr/2024/v26i91269>

Olabanji, S. O., Marquis, Y. A., Adigwe, C. S., Abidemi, A. S., Oladoyinbo, T. O.,

&Olaniyi, O. O. (2024). AI-Driven Cloud Security: Examining the Impact of User Behavior Analysis on Threat Detection. *Asian Journal of Research in Computer Science*, 17(3), 57–74. <https://doi.org/10.9734/ajrcos/2024/v17i3424>

Olabanji, S. O., Olaniyi, O. O., &Olagbaju, O. O. (2024). Leveraging Artificial

Intelligence (AI) and Blockchain for Enhanced Tax Compliance and Revenue Generation in Public Finance. *Asian Journal of Economics, Business and Accounting*, 24(11), 577–587. <https://doi.org/10.9734/ajeba/2024/v24i111577>

Olabanji, S. O., OluwaseunOladejiOlaniyi, O. O., &Olaoye, O. O. (2024). Transforming Tax Compliance with Machine Learning: Reducing Fraud and Enhancing Revenue Collection. *Asian Journal of Economics Business and Accounting*, 24(11), 503–513. <https://doi.org/10.9734/ajeba/2024/v24i111572>

Oladoyinbo, T. O., Olabanji, S. O., Olaniyi, O. O., Adebisi, O. O., Okunleye, O. J., &Alao, A. I. (2024). Exploring the Challenges of Artificial Intelligence in Data Integrity and its Influence on Social Dynamics. *Asian Journal of Advanced Research and Reports*, 18(2), 1–23. <https://doi.org/10.9734/ajarr/2024/v18i2601>

Olaniyi, O. O. (2024). Ballots and Padlocks: Building Digital Trust and Security in Democracy through Information Governance Strategies and Blockchain Technologies. *Asian Journal of Research in Computer Science*, 17(5), 172–189. <https://doi.org/10.9734/ajrcos/2024/v17i5447>

Olaniyi, O. O., Omogoroye, O. O., Olaniyi, F. G., Alao, A. I., &Oladoyinbo, T. O. (2024). CyberFusion Protocols: Strategic Integration of Enterprise Risk Management, ISO 27001, and Mobile Forensics for Advanced Digital Security in the Modern Business Ecosystem. *Journal of Engineering Research and Reports*, 26(6), 32. <https://doi.org/10.9734/JERR/2024/v26i61160>

Olateju, O. O., Okon, S. U., Olaniyi, O. O., Samuel-Okon, A. D., &Asonze, C. U. (2024). Exploring the Concept of Explainable AI and Developing Information Governance Standards for Enhancing Trust and Transparency in Handling Customer Data. *Journal of Engineering Research and Reports*, 26(7), 244–268. <https://doi.org/10.9734/jerr/2024/v26i71206>

Opposs, D. (2020). *The impact of the coronavirus outbreak on exams around the world*.

Ofqual.blog.gov.uk. <https://ofqual.blog.gov.uk/2020/05/22/the-impact-of-the-coronavirus-outbreak-on-exams-around-the-world/>

Reuters. (2025). *BMO Unit to Pay \$40.7 Million in US SEC Settlement Over Misleading Bond Sales*. US News & World Report; U.S. News & World Report.

<https://money.usnews.com/investing/news/articles/2025-01-13/bmo-capital-markets-settles-with-sec-over-bond-desk-supervision>

Reuters. (2025). Italy's regulator blocks Chinese AI app DeepSeek on data protection.

Reuters. <https://www.reuters.com/technology/artificial-intelligence/italys-privacy-watchdog-blocks-chinese-ai-app-deepseek-2025-01-30/>

Ridzuan, N. N., Masri, M., Anshari, M., Fitriyani, N. L., & Syafrudin, M. (2024). AI in the Financial Sector: The Line between Innovation, Regulation and Ethical

Responsibility. *Information*, 15(8), 432. <https://doi.org/10.3390/info15080432>

Sabin, S. (2025). *Shadow AI creates new headaches for company IT teams*. Axios.

<https://www.axios.com/2025/02/04/shadow-ai-cybersecurity-enterprise-software-deepseek>

Sahota, N. (2024). *Shadow AI: The Hidden Risks and Rewards of Unregulated AI Use*.

Neil Sahota. <https://www.neilsahota.com/shadow-ai-the-hidden-risks-and-rewards-of-unregulated-ai-use/>

Salako, A. O., Fabuyi, J. A., Aideyan, N. T., Selesi-Aina, O., Dapo-Oyewole, D. L.,

& Olaniyi, O. O. (2024). Advancing Information Governance in AI-Driven Cloud

Ecosystem: Strategies for Enhancing Data Security and Meeting Regulatory

Compliance. *Asian Journal of Research in Computer Science*, 17(12), 66–88.

<https://doi.org/10.9734/ajrcos/2024/v17i12530>

Samuel-Okon, A. D., Akinola, O. I., Olaniyi, O. O., Olateju, O. O., & Ajayi, S. A. (2024).

Assessing the Effectiveness of Network Security Tools in Mitigating the Impact of

Deepfakes AI on Public Trust in Media. *Archives of Current Research*

International, 24(6), 355–375. <https://doi.org/10.9734/acri/2024/v24i6794>

Schueler, C. (2024). The Shadow AI Dilemma: Balancing Innovation And Data Security In The Workplace. *Forbes*.

<https://www.forbes.com/councils/forbestechcouncil/2024/12/02/the-shadow-ai-dilemma-balancing-innovation-and-data-security-in-the-workplace/>

Seh, A. H., Zarour, M., Alenezi, M., Sarkar, A. K., Agrawal, A., Kumar, R., & Khan, R. A.

(2020). Healthcare data breaches: Insights and implications. *Healthcare*, 8(2),

133. NCBI. <https://doi.org/10.3390/healthcare8020133>

Shandilya, S. K., Datta, A., Kartik, Y., & Nagar, A. (2024). Navigating the Regulatory

Landscape. *EAI/Springer Innovations in Communication and Computing*, 127–

240. https://doi.org/10.1007/978-3-031-53290-0_3

Silor, A. (2024). Navigating the Ethical Landscape: Perspectives on Integrating Artificial

Intelligence in Educational Settings. *SSRN*. <https://doi.org/10.2139/ssrn.4954768>

Val, O. O., Kolade, T. M., Gbadebo, M. O., Selesi-Aina, O., Olateju, O. O., & Olaniyi, O.

O. (2024). Strengthening Cybersecurity Measures for the Defense of Critical

Infrastructure in the United States. *Asian Journal of Research in Computer*

Science, 17(11), 25–45. <https://doi.org/10.9734/ajrcos/2024/v17i11517>

Walsh, B. (2020). *How an AI grading system ignited a national controversy in the*

U.K.Axios. <https://www.axios.com/2020/08/19/england-exams-algorithm-grading>

Walter, Y. (2024). Managing the race to the moon: Global policy and governance in

Artificial Intelligence regulation—A contemporary overview and an analysis of socioeconomic consequences. *Discover Artificial Intelligence*, 4(1).

<https://doi.org/10.1007/s44163-024-00109-4>

Wheeler, K. (2025). *How Trump Scrapping AI Safety Regulations Impacts Global AI.*

Aimagazine.com; Bizclik Media Ltd. <https://aimagazine.com/articles/trump-scraps-ai-risk-rules-what-you-need-to-know>

World Economic Forum. (2025). *Global Cybersecurity Outlook 2025*. World Economic

Forum. <https://www.weforum.org/publications/global-cybersecurity-outlook-2025/>

Yang, E., &Beil, C. (2024). Ensuring data privacy in AI/ML implementation. *New*

Directions for Higher Education, 2024(207). <https://doi.org/10.1002/he.20509>

Zekos, G. (2021). AI Risk Management. *Springer*, 233–288. [https://doi.org/10.1007/978-](https://doi.org/10.1007/978-3-030-64254-9_6)

[3-030-64254-9_6](https://doi.org/10.1007/978-3-030-64254-9_6)